

ITU-T Focus Group Technical Report

(03/2024)

Focus Group on Artificial Intelligence (AI) and Internet of Things (IoT) for Digital Agriculture

FG-AI4A WG-ELR – Ethical, legal and regulatory considerations relating to the use of AI for agriculture: A European perspective

Technical Report ITU-T FG-AI4A

Ethical, legal and regulatory considerations relating to the use of AI for agriculture: A European perspective

Summary

This Technical Report focuses on the ethical considerations arising from the adoption of artificial intelligence (AI) in agricultural production, addressing concerns such as data privacy, transparency, and fairness in algorithmic decision-making. It analyses the legal frameworks applicable to AI in agriculture, emphasizing pertinent European Union (EU) Regulations and Directives relating to data protection, intellectual property rights, and liability. Through this analysis, the report aims to provide a comprehensive understanding of the regulatory landscape surrounding AI technologies in agriculture, offering insights for policymakers and stakeholders.

Keywords

AI, AI Act, digital agriculture, ethical considerations, IoT.

Note

This is an informative ITU-T publication. Mandatory provisions, such as those found in ITU-T Recommendations, are outside the scope of this publication. This publication should only be referenced bibliographically in ITU-T Recommendations.

Acknowledgements

This Technical Report was prepared under the leadership of Mr Sebastian Bosse (Fraunhofer HHI, Germany) and Mr Ramy Fathy (National Telecom Regulatory Authority, Egypt), who serve as the FG-AI4A co-chairs.

It is based on the contributions of various authors who participated in the Focus Group activities. Ms Francesca Hennig-Possenti (John Deere GmbH & Co KG) served as the main editor of this Technical Report. Ms Mythili Menon (FG-AI4A adviser) and Ms Chiara Co (FG-AI4A assistant) served as the FG-AI4A Secretariat.

Change Log

This document contains Version 1.0 of the ITU-T Technical Report on *"Ethical Legal, and regulatory Considerations relating to the use of AI for agriculture: A European Perspective"*, approved at FG-AI4A ninth meeting held in New Delhi on 19 March 2024.

Editor: Francesca Hennig-Possenti
John Deere GmbH & Co KG

E-mail: hennig-possentifrancesca@johndeere.com

© ITU 2025

Some rights reserved. This publication is available under the Creative Commons Attribution-Non Commercial-Share Alike 3.0 IGO licence (CC BY-NC-SA 3.0 IGO; <https://creativecommons.org/licenses/by-nc-sa/3.0/igo>).

If you wish to reuse material from this publication that is attributed to a third party, such as tables, figures or images, it is your responsibility to determine whether permission is needed for that reuse and to obtain permission from the copyright holder. The risk of claims resulting from infringement of any third-party owned material in the publication rests solely with the user.

Table of Contents

		Page
1	Scope.....	1
2	References.....	1
3	Definitions	1
	3.1 Terms defined elsewhere	1
4	Abbreviations and Acronyms	1
5	Introduction.....	2
6	The EU's AI Act	2
	6.1 The AI Act in Europe and Executive Order in the US	3
	6.2 Territorial application of the AI Act.....	4
	6.3 The definition of AI in the AI Regulation: What falls under the concept of AI.....	4
	6.4 Categories of risk in the AI Act.....	5
	6.5 Prohibited AI in the EU's AI Regulation	5
	6.6 High-risk AI.....	6
	6.7 Risk management systems.....	10
	6.8 Data governance requirements	12
	6.9 Concept of error in the AI Regulation	15
	6.10 Sandboxes.....	17
	6.11 Requirements for general purpose AI (e.g., transformer models).....	18
	6.12 Certification.....	18
	6.13 Penalties for non-compliance	19
	6.14 Code of conduct.....	19
7	UN Resolution on AI.....	19
8	AI standards	20
9	Ethical aspects of AI in agriculture	21
10	Conclusion	22
	Bibliography.....	24

Technical Report ITU-T FG-AI4A

Ethical, legal and regulatory considerations relating to the use of AI for agriculture: A European perspective

1 Scope

This Technical Report examines the implications of the European Union (EU) Artificial Intelligence Act (AI Act) and related instruments, including the United Nations (UN) Resolution on Artificial Intelligence, within the context of agriculture.

2 References

- [ITU-T M.3080] Recommendation ITU-T M.3080 (2021), *Framework of artificial intelligence enhanced telecom operation and management (AITOM)*.
- [ITU-T Y.2060] Recommendation ITU-T Y.4000/Y.2060 (2012), *Overview of the Internet of things*.
- [ITU-T Y.4000] Recommendation ITU-T Y.4000 (2012), *Overview of the Internet of things*.
- [ITU-T Y.4450] Recommendation ITU-T Y.4450/Y.2238 (2015), *Overview of Smart Farming based on networks*.
- [ITU-T Y Suppl. 76] Supplement 76 to ITU-T Y series (2023), *ITU-T Y.4000-series – Use cases of Internet of things-based smart agriculture*.

3 Definitions

3.1 Terms defined elsewhere

This Technical Report uses the following terms defined elsewhere:

3.1.1 artificial intelligence (AI) [ITU-T M.3080]: Computerized system that uses cognition to understand information and solve problems.

NOTE 1 – ISO/IEC 2382-28 defines AI as "an interdisciplinary field, usually regarded as a branch of computer science, dealing with models and systems for the performance of functions generally associated with human intelligence, such as reasoning and learning".

NOTE 2 – In computer science AI research is defined as the study of "intelligent agents": any device that perceives its environment and takes actions to achieve its goals.

NOTE 3 – This includes pattern recognition, the application of machine learning and related techniques.

NOTE 4 – Artificial-intelligence is the whole idea and concept of machines being able to carry out tasks in a way that mimics human intelligence and would be considered "smart".

3.1.2 Internet of things (IoT) [ITU-T Y.2060]: A global infrastructure for the information society, enabling advanced services by interconnecting (physical and virtual) things based on existing and evolving interoperable information and communication technologies.

4 Abbreviations and Acronyms

This Technical Report uses the following abbreviations and acronyms:

AI	Artificial Intelligence
FLOP	Floating Point Operation

GMO	Genetically Modified Organism
GPAI	General Purpose AI
ICT	Information and Communication Technology
IoT	Internet of Things

5 Introduction

Ethics and regulations play a pivotal role in shaping the development and implementation of digital agriculture technologies. Ethical considerations ensure that these innovations are deployed responsibly, taking into account factors such as data privacy, cybersecurity, and societal impacts. Regulations provide a framework to safeguard against potential risks and ensure that digital tools are used in a manner that aligns with ethical standards. Regulations may govern the collection and use of agricultural data, ensuring that farmers have control over their information and that data is used in ways that benefit the agricultural community without compromising individual privacy.

Ethical guidelines further promote transparency, fairness, and equity in the adoption of digital technologies, fostering trust among farmers, consumers, and stakeholders. Overall, a harmonious balance between ethics and regulations is essential for the sustainable and ethical advancement of digital agriculture.

6 The EU's AI Act

The AI Act [b-EU AI] is being enacted in Europe as a Regulation and is based on several key principles that reflect the EU's commitment to ethical standards and fundamental rights relating to the conception, development, deployment, and use of artificial intelligence (AI). This approach aims to mitigate potential risks associated with AI technologies, ensuring that they conform to European values and legal standards; safeguarding the safety and fundamental rights of individuals and groups; and fostering the development of secure, trustworthy, and ethical AI. Furthermore, it aims to provide legal certainty to promote the EU's competitiveness in the AI sector.

EU Regulations are binding legislative acts that apply directly in all EU Member States from the date they come into effect, without the need for any national implementing legislation. Regulations have general application, are binding in their entirety, and are directly applicable in all Member States. This means that Regulations have the power to create rights and obligations for individuals and entities across the EU simultaneously.

The AI Act is applicable from the date of its publication, with specific deadlines for parts prescribed in the Regulation itself.

The AI Act is conceived as a proactive or ex ante compliance Regulation. Its primary aim is to establish guidelines and requirements that must be adhered to before a product or a service is placed on the market, thereby ensuring that AI-driven technologies are developed and utilized in a manner that is safe, ethical, and aligned with fundamental rights.

EU Directives, on the other hand, are legislative acts that set out goals that all EU countries must achieve. However, unlike Regulations, Directives do not prescribe how these goals are to be achieved. This allows Member States the flexibility to adapt the Directive to their own legal systems and national circumstances, through the enactment of domestic legislation within a set deadline.

An example of an EU Directive is the planned Artificial Intelligence Liability Act, which sets ambitious targets for all EU countries, to regulate the liability arising from the deployment and use of AI systems in Europe.

The planned Artificial Intelligence Liability Act serves as a reactive legal mechanism designed to provide a legal recourse for scenarios in which damage or harm results from the use of AI technologies.

The AI Act focuses on preventive measures and standards for the responsible innovation and application of AI, while the Artificial Intelligence Liability Act concentrates on the attribution of responsibility and the resolution of legal disputes arising post-incident, when damage has already occurred. This Directive will most likely be postponed until after the European election in the spring of 2024. In the meantime, the AI Act appears to be subject to the product liability rules provided for digital products such as software.

These legislative instruments aim to provide a legal framework for AI, managing the life cycle of AI technologies, from development and deployment on the market to the aftermath of their application and/or use.

The AI Act aims to establish a comprehensive regulatory framework for AI systems, to ensure safety, transparency, and respect for fundamental rights while fostering innovation.

6.1 The AI Act in Europe and Executive Order in the US

Comparing the AI Act to US President Joe Biden's Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, published on 30 October 2023, both legislations prioritize safety, ethical considerations, and responsible innovation. However, the AI Act is more regulatory, categorizing AI applications by technology-related risks and setting specific compliance requirements.

President Biden's Executive Order emphasizes principles and a collaborative, multi-stakeholder process to govern AI development and use, empowering dedicated institutions to provide sectoral regulations. Both legislations aim to balance innovation with safeguards against potential harm; however, Biden's Executive Order introduces the concept of the industry's self-responsible approach to innovation.

The Executive Order involves engaging with a broad range of stakeholders to identify and address ethical, legal, societal, and policy implications early in the innovation process, putting the emphasis on responsible innovation given the broader societal impact and potential risks associated with AI use.

In this sense, the Executive Order does not create categories of technologies that are, by definition, high risk but, instead, focuses more on the actual risk evaluation in consideration of their use. Thus, systems that may fall under the high-risk classification under the rules of the AI Act may not necessarily be considered risky under the Executive Order, depending on the actual AI usage environment. This is especially relevant for agricultural applications, where, compared to other sectors, the risks associated with AI appear more limited due to the nature of the operations being conducted in rural or secluded areas with minimal human presence.

Both legislative measures aim to balance innovation with safeguards against potential harm, with the AI Act focusing on preventive measures and standards for responsible innovation and the Executive Order involving early engagement with stakeholders to address ethical, legal, societal, and policy implications associated with AI. The differences in approach reflect the distinct regulatory and governance strategies of the EU and the US in managing the development and use of AI technologies.

6.2 Territorial application of the AI Act

The AI Act "applies to providers placing on the market or putting into service AI systems or placing on the market general-purpose AI models in the Union, irrespective of whether those providers are established or who are located within the Union or in a third country"¹.

In this regard, the Regulation applies to providers of AI from the moment the services are offered in the EU or to EU citizens, regardless of the provider/producer establishment or geographical location, whether within the Union or in a third country. For example, a company or institution incorporated or established in Asia or the US may be subject to the AI Act if the services are offered in Europe.

Additionally, results used in the EU, if produced by AI, may require providers to be compliant with the AI Act, as the Regulation states: "providers and deployers of AI systems that have their place of establishment or who are located in a third country, where the output produced by the system is used in the Union."

The Regulation's wording seems to extend the application of the law to services provided outside the EU but whose results are sold or provided to European stakeholders. For example, an AI company located in India may compile the software to be executed automatically by a robot in Europe, or an AI credit-scoring service offered in the US may provide credit calculation results on EU customers to be used in Europe.

Producer obligations to comply with the law may be transferred to importers and distributors of AI systems if the producer company is not located in Europe. The same applies to products that embed the AI system or place the AI system on the market under a local producer name (rebranded products).

The definition of the territorial scope and application provides a very broad application of the Regulation, which extends way beyond the territory of the European Union.

6.3 The definition of AI in the AI Regulation: What falls under the concept of AI

The definition in Art. 3 of the AI Act is as follows: "An AI system is a machine-based system designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments"².

The AI Act defines an "AI system" as a machine-based system, encompassing software, algorithms, models, and training with various characteristics, and allows for new methods to be included, such as pre-trained models applied in the training of different narrow AIs instead of training an AI system from scratch.

The Act sets further criteria for a system to be defined as AI, including the requirement for the system to be "designed to operate with varying levels of autonomy". This reference has been discussed at length in the automotive sector, providing a description of different autonomy stages determined by the different stages of human interaction with the machine functioning and decision patterns. This level applies to autonomous guidance and also to virtual decision processes, such as assessing the maturity of fruits, determining the quality of soil, or steering autonomous operations.

To differentiate autonomous systems from automated systems, the AI system is required to possibly "exhibit adaptiveness after deployment". The use of the term "may" leaves open questions about the interpretation of this requirement, which may extend to the qualification of AI for systems that are not adaptive, challenging the understanding of deterministic systems in a classical interpretation of the term.

¹ Art. 2 AI Act.

² Art. 3 (1) of the AI Act Draft – definition of AI.

As part of the AI qualifying criteria, the AI Act requires the system, "for explicit or implicit objectives", to infer, "from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions," implying that the AI system is designed to achieve certain goals or objectives, whether clearly defined by humans or not directly stated.

The term "infer" implies that the system is designed to modify, adjust, and act upon the input data to achieve its objectives. This suggests that a system may be qualified as AI if designed to derive insight or make decisions based on the input it receives, particularly if the system is programmed or trained on methodologies to process input data and generate various forms of output.

The definition also encompasses software systems that demonstrate a capacity for human-like analysis, decision-making, and learning, including machine learning models, logic and knowledge-based systems, and statistical approaches such as Bayesian estimations and search and optimization methods. The aim of the legislator was to provide an open and adaptable definition that remains relevant as new forms of AI emerge, avoiding the pitfall of fast technological obsolescence. However, this broad definition also poses challenges, particularly in delineating the boundaries of what constitutes AI within a legal context, potentially leading to ambiguities in regulatory compliance and enforcement.

Furthermore, the definition of AI points out the capability of the system to "influence physical or virtual environments". However, the legislator seems to consider this descriptive element only as optional, pointing out, with the word "can", that the capability of influencing the environment is not mandatory in the identification of AI.

In summary, the definition of AI continues to challenge interpreters due to its general formulation.

6.4 Categories of risk in the AI Act

The AI Act defines risk as "the combination of the probability of an occurrence of harm and the severity of that harm".

In legal terms, the concept of "risk" encompasses the likelihood and potential severity of harm or adverse effects arising from the AI system. It quantifies the possibility of negative outcomes and is central to evaluating compliance, safety, and accountability across various domains.

In the context of law, especially within regulatory frameworks such as the AI Act, risk assessment involves identifying, evaluating, and managing the potential for harm to individuals, groups, or society at large.

This process is crucial for developing strategies to mitigate or prevent harm, ensuring that actions and innovations align with legal standards, ethical principles, and societal values. The categorization of risks, particularly in the field of AI, guides the application of regulatory measures, ranging from minimal oversight for low-risk scenarios to strict controls and prohibitions for high-risk applications.

6.5 Prohibited AI in the EU's AI Regulation

The EU AI Act includes, in Art. 5, several provisions regarding prohibited AI practices. It clarifies and expands on prohibitions concerning the use by law enforcement of real-time biometric identification in publicly accessible spaces, with some exceptions that are subject to stringent safeguards, monitoring, and limited reporting at the EU level. Other key prohibitions involve the scraping of facial images for creating facial recognition databases, emotion recognition in workplaces and educational settings (with safety and medical exceptions), and certain forms of biometric categorization and predictive policing.

Furthermore, the AI Act prohibits AI systems designed to manipulate human behaviour, exploiting vulnerabilities of specific groups or the classification of individuals based on social scoring or other indicators. The application of those principles has already fired discussion about the application of the Regulation to existing credit scoring, pointing out the need to discipline the sector more.

These prohibitions are designed with a clear intent to balance the need for technological innovation and public safety with the protection of fundamental rights and freedoms, addressing ethical concerns such as privacy, autonomy, and discrimination.

While AI is still evolving, some examples of prohibited AI may include:

- 1) Placing on the market or using an AI system that deploys subliminal techniques or purposefully manipulative or deceptive techniques to materially distort a person's behaviour and impair their ability to make informed decisions; AI systems that exploit the vulnerabilities of a specific group of persons, such as those due to their age, disability, or social or economic situation; biometric categorization systems to deduce sensitive personal information, such as race, political opinions, or sexual orientation; AI systems to evaluate or classify individuals based on their social behaviour or personality characteristics, leading to unjustified detrimental treatment in social contexts, such as, for example, an AI system that assigns social scores to individuals and leads to unjustified exclusion from social activities or opportunities.
- 2) Real-time use of remote biometric identification systems in publicly accessible spaces for law enforcement purposes is prohibited except in specific circumstances, such as preventing imminent threats or locating missing persons; AI systems for making risk assessments of individuals solely based on profiling or personality traits, to predict criminal behaviour such as, for example, an AI system that predicts the likelihood of an individual committing a crime based solely on their social media activity and personality traits.
- 3) AI systems to create or expand facial recognition databases through untargeted scraping of facial images from the Internet or CCTV footage such as, for example, an AI system that scrapes or collects facial images from social media without individuals' consent, to build a facial recognition database; AI systems to infer emotions of individuals in workplaces and educational institutions, unless it is for medical or safety reasons such as, for example, an AI system used in workplaces to analyse employees' emotions without their consent for performance evaluations.

A different set of prohibitions may originate from the use of AI to develop knowledge or applications that are prohibited by other legislations. The relevant prohibition is not to be found in the AI Act but more in the specific legislation. For example (but not only) genetic manipulation, particularly in the context of genetically modified organisms (GMOs), is subject to strict regulations in Europe due to concerns about potential environmental and human health risks. The EU has established a comprehensive legal framework to govern the approval, marketing, and use of GMOs, and this framework reflects the precautionary principle, which is a core element of EU environmental and public health policy³. The use of AI is not subject to a specific prohibition but may be considered a tool for achieving a prohibited outcome already disciplined by other legislations.

6.6 High-risk AI

Article 6 of the AI Act defines a high-risk AI system based on specific classification rules. Article 6 states that one of the conditions that define whether an AI system is high risk is if it is "intended to be used as a safety component or is itself a product".

³ The building blocks of the GMO legislation provide the basic rules, recommendations and guidelines on GMO use in Europe: the Directive 2001/18/EC on the deliberate release of GMOs into the environment; the Regulation (EC) 1829/2003 on genetically modified food and feed; the Directive (EU) 2015/412 amending Directive 2001/18/EC regarding the possibility for Member States to restrict or prohibit the cultivation of GMOs in their territory; the Regulation (EC) 1831/2003 concerning the traceability and labelling of GMOs and the traceability of food and feed products produced from GMOs; the Directive 2009/41/EC on contained use of genetically modified micro-organisms; and the Regulation (EC) 1946/2003 on transboundary movements of GMOs.

The provision seems to include AI standalone applications and systems such as desktop applications or apps (e.g., including, but not only, AI remote equipment control) as well as AI systems embedded in the agricultural and/or forestry machine (e.g., including, but not only, AI embedded in robots).

Art. 2 of the AI Act defines "safety component" as "a component of a product or of a system which fulfils a safety function for that product or system, or the failure or malfunctioning of which endangers the health and safety of persons or property;".

This definition seems to differ from the definition of safety component in the Machinery Regulation, which defines "safety component" as a "physical or digital component, including software, of a product within the scope of (the Machinery) Regulation, which is designed or intended to fulfil a safety function and which is independently placed on the market, the failure or malfunction of which endanger the safety of persons, but which is not necessary in order for that product to function or for which normal components may be substituted in order for that product to function".

The AI Act introducing the endangering of property seems to extend the application of Art. 6 further than the Machinery Regulation that focuses on the safety of persons extending the compliance protection to property damage usually protected by ex post legislation (after the damage is done) through civil and tort law provisions. The extensive compliance protection provided by the definition of safety component in the AI Act applies without the need for actual damage, even if no person is at risk of being endangered by the system.

In addition, while the Machinery Regulation provides a definition of safety function "as a function that serves to fulfil a protective measure designed to eliminate, or, if that is not possible, to reduce, a risk, which, if it fails, could result in an increase of that risk", the AI Act does not provide any definition of safety function in the current draft, leaving it open to interpretation and underlining the requirement for a more harmonized interpretation.

Furthermore, for the classification of high risk, Art. 6 provides the additional cumulative condition that the system must be already "covered by the EU legislation listed in Annex I (former Annex II) of the AI Act". The Annex lists several Regulations and Directives applicable to different sectors, including among others, the Machinery Directive (soon to be substituted by the Machinery Regulation) and the Agricultural and Forestry Equipment Regulation.

For a product to be classified as high risk, the AI Act submits to the NLF to mandate if the product needs to undergo a third-party conformity assessment. Additionally, Art. 6 provides, as a requirement for the AI to be designated as high risk, that the AI system is itself a product (for example, an app or service) or is intended to be used as a safety component of a product (autonomy component).

Examples of machinery and applications that may be used in farm and forestry environments included in the application of Annex I are as follows: different kinds of agricultural robots and non-road machinery⁴, tractors and forestry machines⁵ but also watercraft vehicles⁶ and marine equipment⁷,

⁴ Directive 2006/42/EC of the European Parliament and of the Council of 17 May 2006 on machinery, and amending Directive 95/16/EC [as repealed by the Machinery Regulation].

⁵ Regulation (EU) No 167/2013 of the European Parliament and of the Council of 5 February 2013 on the approval and market surveillance of agricultural and forestry vehicles.

⁶ Directive 2013/53/EU of the European Parliament and of the Council of 20 November 2013 on recreational craft and personal watercraft.

⁷ Directive 2014/90/EU of the European Parliament and of the Council of 23 July 2014 on marine equipment and repealing Council Directive 96/98/EC.

lifts⁸, equipment or protective systems intended for use in potentially explosive atmospheres⁹, radio equipment¹⁰, pressure equipment¹¹, appliances burning gaseous fuels¹², and agricultural drones¹³.

The list of equipment and machines that may be subject to inclusion under the high-risk requirements also includes products that are not used in agriculture or forestry such as toys, medical equipment and in vitro diagnostics.

6.6.1 High-risk AI

Art. 6 of the AI Act defines high-risk AI systems based on specific classification rules. Article 6 states that one of the conditions that define whether an AI system is high risk is if it is "intended to be used as a safety component or is itself a product".

The provision seems to include AI standalone applications and systems, such as desktop applications or apps (e.g., including, but not only, AI remote equipment control) as well as AI systems embedded in the agricultural and/or forestry machinery (e.g., including, but not only, AI embedded in robots).

Art. 2 of the AI Act defines "safety component" as "a component of a product or of a system which fulfils a safety function for that product or system, or the failure or malfunctioning of which endangers the health and safety of persons or property;".

This definition seems to differ from the definition of safety component in the Machinery Regulation that defines "safety component" as a "physical or digital component, including software, of a product within the scope of (the Machinery) Regulation, which is designed or intended to fulfil a safety function and which is independently placed on the market, the failure or malfunction of which endanger the safety of persons, but which is not necessary in order for that product to function or for which normal components may be substituted in order for that product to function".

The AI Act, introducing the endangering of property, seems to extend the application of Art. 6 further than the Machinery Regulation, which focuses on the safety of persons, extending the compliance protection to property damage that is usually protected by ex post legislations (after damage is done) with civil and tort law provisions. The extensive compliance protection provided by the definition of safety component in the AI Act applies without the need for actual damage, even if no person is at risk of being endangered by the system.

In addition, while the Machinery Regulation provides a definition of safety function "as a function that serves to fulfil a protective measure designed to eliminate, or, if that is not possible, to reduce, a risk, which, if it fails, could result in an increase of that risk", the AI Act does not provide any definition of safety function in the current draft, leaving it open to interpretation and underlining the requirement for a more harmonized interpretation.

Furthermore, for the classification of high risk, Art. 6 provides the additional cumulative condition that the system must be already "covered by the EU legislation listed in Annex I (former Annex II) of the AI Act".

⁸ Directive 2014/33/EU of the European Parliament and of the Council of 26 February 2014 on the harmonisation of the laws of the Member States relating to lifts and safety components for lifts.

⁹ Directive 2014/34/EU of the European Parliament and of the Council of 26 February 2014.

¹⁰ Directive 2014/53/EU of the European Parliament and of the Council of 16 April 2014.

¹¹ Directive 2014/68/EU of the European Parliament and of the Council of 15 May 2014.

¹² Regulation (EU) 2016/426 of the European Parliament and of the Council of 9 March 2016.

¹³ Regulation (EC) No 300/2008 of the European Parliament and of the Council of 11 March 2008; Regulation (EU) 2018/1139 of the European Parliament and of the Council of 4 July 2018 on common rules in the field of civil aviation and establishing a European Union Aviation Safety Agency, in so far as the design, production and placing on the market of aircrafts, where it concerns unmanned aircraft and their engines, propellers, parts and equipment to control them remotely.

The Annex lists several Regulations and Directives applicable to different sectors, including, among others, the Machinery Directive (soon to be substituted by the Machinery Regulation) and the Agricultural and Forestry Equipment Regulation. For a product to be classified as high risk, the AI Act submits to the NLF to mandate if the product needs to undergo a third-party conformity assessment. Additionally, Art. 6 provides, as a requirement for the AI system to be designated as high risk, that the AI system is itself a product (for example, an app or service) or is intended to be used as a safety component of a product (autonomy component).

Examples of machinery and applications that may be used in farm and forestry environment included in the application of Annex I are as follows: different kinds of agricultural robots and non-road machinery¹⁴, tractors and forestry machines¹⁵ but also watercraft vehicles¹⁶ and marine equipment¹⁷, lifts¹⁸, equipment or protective systems intended for use in potentially explosive atmospheres¹⁹, radio equipment²⁰, pressure equipment²¹, appliances burning gaseous fuels²², and agricultural drones²³.

The list of equipment and machines that may be subject to inclusion under high-risk requirements also includes products that are not used in agriculture or forestry, such as toys, medical equipment, and in vitro diagnostics.

6.6.2 High-risk in critical infrastructures

Art. 6 (2) provides the classification for high risk for AI operating in critical infrastructures. The definition of critical infrastructure is not provided by the nature of the activity itself but is the result of a closed list enumerated by the AI Act on the basis of its relevance in relation to the protection of human rights or considering the importance of a specific sector for societal well-being.

A second category includes road traffic, the supply of essential services such as water, gas, heating, and electricity. The classification as critical infrastructure relies on the impact that disruption may have on the conduct of social and economic activities on a large scale. Agriculture, which was considered as critical infrastructure under the Corona provision, allowing for the exclusion of agricultural production from the restrictions issued in 2020, has not yet been considered as critical infrastructure under this category. Examples may be an AI system powering a farm energy grid, a large-scale water management system, or a local energy power plant (pellets or biogas) selling energy to the community grid.

¹⁴ Directive 2006/42/EC of the European Parliament and of the Council of 17 May 2006 on machinery, and amending Directive 95/16/EC [as repealed by the Machinery Regulation].

¹⁵ Regulation (EU) No 167/2013 of the European Parliament and of the Council of 5 February 2013 on the approval and market surveillance of agricultural and forestry vehicles.

¹⁶ Directive 2013/53/EU of the European Parliament and of the Council of 20 November 2013 on recreational craft and personal watercraft.

¹⁷ Directive 2014/90/EU of the European Parliament and of the Council of 23 July 2014 on marine equipment and repealing Council Directive 96/98/EC.

¹⁸ Directive 2014/33/EU of the European Parliament and of the Council of 26 February 2014 on the harmonisation of the laws of the Member States relating to lifts and safety components for lifts.

¹⁹ Directive 2014/34/EU of the European Parliament and of the Council of 26 February 2014.

²⁰ Directive 2014/53/EU of the European Parliament and of the Council of 16 April 2014.

²¹ Directive 2014/68/EU of the European Parliament and of the Council of 15 May 2014.

²² Regulation (EU) 2016/426 of the European Parliament and of the Council of 9 March 2016.

²³ Regulation (EC) No 300/2008 of the European Parliament and of the Council of 11 March 2008; Regulation (EU) 2018/1139 of the European Parliament and of the Council of 4 July 2018 on common rules in the field of civil aviation and establishing a European Union Aviation Safety Agency, in so far as the design, production and placing on the market of aircrafts, where it concerns unmanned aircraft and their engines, propellers, parts and equipment to control them remotely.

The main effect of powering an AI system for a business operating under one or more of those categories is that all AI systems, independent of their features, are, by default, considered to be high risk and required to fulfil the requirements applicable to high-risk AI.

To mitigate the effects of the Regulation on simpler AI that are used in critical infrastructure settings (e.g., an intelligent beverage distributor in a power plant environment), EU institutions have opted for a special regimen for low-risk AI in critical infrastructures.

In this sense, AI systems will not be classified as high risk if they do not present a significant risk of harm to the health, safety, or fundamental rights of individuals, including AI not substantially impacting decision-making outcomes. This exemption applies, for example, if the AI system is designed to carry out a specific procedural task and/or the AI system is intended to enhance the outcome of a previously completed human activity. Furthermore, an AI system is not considered as high risk if it is intended to identify decision-making patterns or deviations from previous decision-making patterns, and is not intended to replace or influence the previously completed human assessment without appropriate human review. This is also the case if an AI system is intended to perform a preparatory task.

In those cases, the legislators allow manufacturers and providers to issue a self-certification to exclude the AI from the high-risk category. However, those certifications constitute a reversible declaration, potentially susceptible to be dismissed by the authorities, resulting in the need to fulfil the requirements provided for high-risk AI.

6.6.2.1 Other high-risk AI

Beside critical infrastructures, the AI Act identifies further AI applications that are to be considered as high risk because of the sensitive field of use or the nature of the AI. The listing provides the classification as high risk for several AI applications, such as biometrics, HR, credit scoring and insurance, access to education and social services, provision of medical and/or emergency services, law enforcement, justice, criminal prevention, or border protection.

6.6.2.2 Requirements for high-risk AI under the AI Act

The AI Act categorizes certain AI systems as high risk due to their significant implications for health, safety, and fundamental rights. For these high-risk AI systems, the Act sets forth stringent requirements to ensure their safe, transparent, and accountable use. The requirements for high-risk AI systems include, among others, the provision of risk management systems, data governance procedures, human oversight, technical documentation, transparency requirements and robustness.

6.7 Risk management systems

High-risk AI management systems must identify, analyse, and forecast risks, providing mitigation strategies that ensure a controlled environment for AI deployment and use.

Developers and deployers of high-risk AI systems must implement a comprehensive risk management system, continually assessing and mitigating risks associated with the AI system's deployment and use.

Art. 8 provides that high-risk AI systems must comply with the requirements specified in the AI Act, taking into account their intended purpose and the current state of the art in AI and AI-related technologies. Providers of products containing AI systems are responsible for ensuring full compliance with all applicable requirements. To streamline processes and minimize additional burdens, providers have the option to integrate necessary testing, reporting processes, and documentation with existing documentation and procedures required under Union harmonization legislation.

The requirement of a risk management system has also been identified by standardization organizations²⁴.

The AI Act mandates the establishment of a structured risk management system to identify and analyse the known and reasonably foreseeable risks associated with the use of high-risk AI systems. This risk management system must consider various aspects of AI programming. For instance, in the case of an autonomous vehicle AI system, a known risk could be the failure to detect and appropriately respond to unexpected hazards, leading to accidents and injuries. Similarly, in an AI system used for automated weed control with lasers or fire devices, a foreseeable risk could be the unintended start of a field fire.

In the context of an AI-powered agricultural drone system, a potential risk could be the invasion of privacy if the drone's surveillance capabilities capture images of individuals without their consent. Likewise, in an AI-powered agricultural robot system, a possible risk could be the unintended physical interaction with humans or animals, resulting in injuries due to a lack of accurate obstacle detection.

In an AI-driven autonomous grain sorting and processing system, a known risk could be the misclassification of grains, leading to a contamination of food products with allergens or toxins, posing a health risk to consumers.

These examples demonstrate the importance of identifying and analysing risks related to high-risk AI systems, taking into account their potential impact on the health, safety, and fundamental rights of individuals when used in accordance with their intended purpose. The "estimation and evaluation of the risks that may emerge when the high-risk AI system is used in accordance with its intended purpose and under conditions of reasonably foreseeable misuse"²⁵ is to be interpreted as the holistic analysis of the risks, considering possible scenarios where the AI system may be subject to potential misuse scenarios that can be reasonably anticipated based on the intended purpose of the AI system and the context of its use, as well as taking into account the severity of the risk and the probability of the negative event occurring.

Examples of a reasonably foreseeable misuse could be²⁶ in an AI-driven autonomous agricultural drone system intended for crop monitoring, where a foreseeable misuse could be the unauthorized use of the drone for surveillance of neighbouring properties, infringing on privacy rights and data protection regulations, or in an AI-powered autonomous harvesting system, where a foreseeable misuse could be the deliberate reprogramming of the system to operate at unsafe speeds or with modified harvesting patterns, compromising worker safety and equipment integrity.

The risk management system is required to provide an "evaluation of other possibly arising risks based on the analysis of data gathered from the post-market monitoring system"²⁷, including in this analysis risks arising from the further development of the AI on the field or during its use.

The AI system's risk management shall be capable of identifying such risks and provide mitigation strategies for a better product design. "The risks (...) shall concern only those which may be reasonably mitigated or eliminated through the development or design of the high-risk AI system, or the provision of adequate technical information."

The requirement seems to point to the control of different risks that may occur in the development phase, requiring provision for the possibility of a design review procedure by analysing the AI input elements, such as, but not only, the chosen methods, algorithms, data libraries, pre-training

²⁴ NIST AI 100-1.

²⁵ Art. 9 AI Act.

²⁶ Those examples are not actual examples but only theoretical possibilities used to provide a possible application.

²⁷ Art. 9 AI Act.

applications, models and training methods. While it is not possible or difficult in complex systems (yet) to clearly establish a nod activation causality, the requirement points to the design elements that are under the control of the development teams.

Furthermore "the risk management measures (...) shall give due consideration to the effects and possible interaction resulting from the combined application of the requirements (...), with a view to minimising risks more effectively while achieving an appropriate balance in implementing the measures (...)". AI risks are particularly difficult to assess, as they may require an analysis of the interaction of multiple causes and effects that may lead to the damaging event in order to prevent and correct dangerous or negative outcomes, or even providing additional measures to mitigate possible arising risks. In this regard, the AI Act states that "The risk management measures (...) shall be such that relevant residual risk associated with each hazard as well as the overall residual risk of the high-risk AI systems is judged to be acceptable". This could be the case, for example, in requiring an additional measure, such as an access control where autonomous devices are operating in a greenhouse or by providing appropriate training to the operator.

Furthermore, the legislator is considering the impossibility of forecasting all possible variables that lead to a negative AI output. In this respect, the legislator requires a system capable of the "implementation of adequate mitigation and control measures addressing risks that cannot be eliminated."

To minimize or mitigate risks associated with the use of high-risk AI systems, careful attention must be given to the technical expertise, practical experience, education, and training expected from the deployer, as well as the likely context in which the system is intended to be used. High-risk AI systems must undergo extensive testing (virtual, labour and real world) to identify the most suitable and specific risk management measures. This testing is essential to ensure the consistent performance of high-risk AI systems for their intended purpose and their compliance with the AI Act.

The testing of high-risk AI systems should be conducted as necessary, at any stage throughout the development process and, in any case, before the product is made available on the market or put into service. Testing should be carried out against predefined metrics and probabilistic thresholds that align with the intended purpose of the high-risk AI system. Standards may be developed to assess the metrics and thresholds to be applied to specific applications, such as (but not only) the required percentage of error-free image recognition for autonomous machines when the system is required to recognize humans.

When implementing the risk management system, providers must consider whether the high-risk AI system, given its intended purpose, is likely to have adverse effects on individuals requiring a particular ethical analysis, such as (but not only) assessing the libraries' data to prevent algorithmic discrimination (e.g., the system do not achieve the required thresholds in recognizing persons of specific ethnicities in an operating environment).

6.8 Data governance requirements

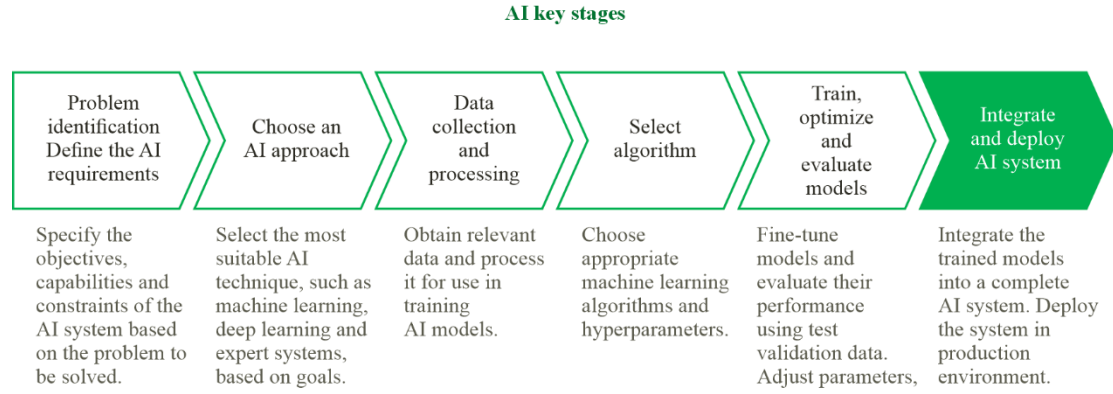
High-risk AI systems must utilize high-quality datasets to train, validate, and test the AI. This ensures the system's performance is reliable and free from biases. Proper data governance mechanisms must be in place to handle data securely and ethically.

AI systems, such as those used in agricultural machinery, may undergo training on datasets that are subject to change over time, sometimes in unexpected and significant ways. This dynamic nature of data can impact the functionality and reliability of AI systems, posing challenges in understanding and maintaining their trustworthiness. Additionally, the complexity of AI systems and their deployment contexts can make it challenging to detect and address failures effectively. The risks and benefits associated with AI systems stem from the interplay of technical aspects and variables, including how the system is utilized, its interactions with other AI systems, the operators involved, and the environment in which it is deployed.

Standardization frameworks, such as NIST 100-1²⁸, approach the analysis of risks related to AI by assessing the trustworthiness of the system, while the AI Act aims to achieve the same objectives by providing detailed technical requirements for AI, in Art. 10 dedicated to the data governance of AI.

While Art. 9 provides the requirement for an overall risk management system, Art. 10 goes into the details of the governance of AI systems. Training, validation, and test datasets shall be subject to data management methods that are appropriate for the intended purpose of the AI system.

The paragraph analyses the different AI development stages, depicted in Figure 1, considering the design requirements in the different stages:



Francesca Henning-Possenti, January 2024

FGAI4A-ELR(25)

Figure 1 – AI key stages

Art. 10 of the AI Act addresses the "presentation of the relevant design decisions"²⁹. This provision postulates the requirement to assess and document the design decisions made throughout the development of the AI system. This should encompass the selection of algorithms, models, and architecture, as well as the rationale behind these choices. Additionally, it should provide insights into how these decisions align with the intended purpose of the AI system and how they contribute to risk management and compliance with regulatory requirements.

Furthermore, the Article requires a "description of the data collection processes and origin of the data and, in the case of personal data, the original purpose of the data collection". This involves providing a comprehensive description of the data collection processes, including the sources of the data and the methods used to gather, store, and manage it. For personal data, the original purpose of the data collection should be clearly articulated, highlighting the ethical and legal considerations related to data privacy and consent.

In the requirement of "presentation of the relevant processing operations in data preparation, such as annotation, labelling, cleansing, updating, enrichment and aggregation", the legislator addresses the various processing operations involved in data preparation, such as annotation, labelling, cleansing, updating, enrichment, and aggregation independently from the chosen learning methods.

In substance, the legislation may require providers to analyse their data cleaning methods, including the handling of missing values, correcting, and threading. It may target the library integration from multiple sources. This may play a particular role in the case of partially or fully pre-trained systems. It may require the description of annotation and labelling procedures (if any) and include a description of annotation and labelling procedures purchased from third-party providers. Furthermore, it may require the provider to describe data normalization procedures and feature attribution as well as data

²⁸ [Artificial Intelligence Risk Management Framework \(AI RMF 1.0\) \(nist.gov\)](https://nist.gov/artificial-intelligence-risk-management-framework-ai-rmf-1.0)

²⁹ Art. 10 of the AI Act.

reduction or cluster analysis methods or other aggregation procedures. Feature scaling and data variables may need to be displayed as well as procedures to handle data imbalances or data representativeness in distribution disparities. Overfitting detection and/or avoidance procedures may play a role in the definition of the required library. Data encoding/transformation procedures may require to be explained, as well as the procedures creating subsets of data.

All those requirements aim to provide transparency in the preliminary steps to AI data analysis, ensuring data quality, integrity, and relevance, and providing a key to addressing potential biases or errors in the dataset, including (but not only), the absence of appropriate databases containing diverse images for human image recognition processes.

In the "Description of the formulation of assumptions, in particular with regard to the information/processes to be measured and presented with the data", the legislator requires AI system providers to deliver a clear and transparent account of the assumptions made during the formulation of the AI model, in view of the scope to be achieved. This includes assumptions related to the information and processes being measured and presented with the data, and how these assumptions impact the accuracy, fairness, and reliability of the AI system. For example, in the context of AI in agriculture and the development of an AI model for crop yield prediction in precision agriculture, the assumption is made that the historical crop yield data collected from a specific region is representative of future crop yield patterns. The AI model assumes that the environmental conditions, soil quality, and other relevant factors affecting crop yield will remain consistent or follow predictable trends over time (e.g., global warming).

Additionally, the model assumes that the input data (such as weather parameters, soil data, and crop health indicators) used for training the AI model accurately captures the complex relationships and interactions that influence crop yield. A clear and transparent account of these assumptions is crucial for understanding how they influence the accuracy, fairness, and reliability of the AI system. This includes assessing the potential impact of changing environmental conditions, evolving farming practices, and unforeseen factors that may affect the predictive capability of the AI model. By clearly articulating these assumptions and their potential impact, providers can demonstrate their commitment to transparency and ethical responsibility in the development and deployment of AI systems in agricultural machinery.

"Evaluation of the availability, quantity and suitability of the required data sets" may require evaluations to be conducted to assess the availability, quantity, and suitability of the datasets used in the development and training of the AI system, to avoid aberrant results due to the reduced diversity or poor quality of the database. In this regard, the analysis should include an examination of data quality, representativeness, and diversity, as well as an assessment of the datasets' sufficiency to support the intended use of the AI system.

Furthermore, the provider is required to provide "assessment methods with regard to possible errors/biases that may affect the health and safety of persons, have a negative impact on fundamental rights or lead to discrimination prohibited by Union law, in particular where the data results influence the inputs for future transactions".

The requirements can be interpreted as providing the obligation to conduct an ethical impact assessment to evaluate the potential societal and ethical implications of the AI system. This involves considering the impact of the system's decisions on different stakeholders and identifying potential biases that may lead to adverse effects. An ethical analysis should be conducted to assess the potential errors, biases, and ethical implications of the data used in the AI system. This is the case, for example, if the providers acknowledge a lack in the database that leads to an algorithmic discrimination that results in a failure of the AI system to recognize humans of different ethnicities and features.

This requirement can also require an assessment to evaluate the impact on the health, safety, and fundamental rights of individuals, and identify any discriminatory or biased outcomes that may contravene Union law. However, this analysis may be particularly challenging in agriculture, as AI

systems may touch the interests of different stakeholders, e.g., by tackling labour shortage, but at the same time, eliminating the need for human labour in the sector or requiring a redirection/requalification of employees and seasonal workers.

"Presentation of the appropriate measures to detect, avoid and mitigate possible errors/biases" involves presenting a comprehensive set of measures designed to detect, avoid, and mitigate potential errors and biases in the data and the AI system. This includes outlining preprocessing techniques, fairness-aware algorithms, and bias mitigation strategies implemented to ensure ethical and unbiased AI outcomes. This involves, for example, assessing reasons for AI failures in recognizing humans in specific settings that requires integrating databases or integrating additional sensors (e.g., lidar or heat sensors) to achieve desired outcomes.

Developing "[m]ethods of determining relevant data gaps or deficiencies, and the processes/methods for resolving bias/errors" involves detailing the methods used to identify and address potential data gaps or deficiencies, including biases and errors in the training data. The processes and techniques employed to detect and resolve data-related issues that may impact the fairness and accuracy of the AI system should be outlined.

To fulfil the requirement of "[a]cquiring relevance of the training, validation, and test datasets (relevant and representative) with regard to the intended use", the relevance of the training, validation, and test datasets should be assessed and validated regarding their representativeness, diversity, and relevance to the intended use of the AI system. This includes ensuring that the datasets are comprehensive, unbiased, and suitable for evaluating the performance and generalization of the AI model.

6.9 Concept of error in the AI Regulation

An error-free AI system is intended to produce accurate outputs that reflect the underlying patterns and relationships within the processed data. This involves making correct classifications, predictions, and identifications based on the input data and the desired outcomes. However, as a statistical system, achieving a completely error-free state is conceptually almost impossible. While AI systems can be designed to improve error finding and correction processes, they may never reach a completely error-free state. The requirement of a possible error-free system should be interpreted with a degree of approximation, requiring the provider to make the best effort to avoid aberrant results and to correct detected errors.

The concept of appropriate statistical properties for data is also open to discussion, as it involves the characteristics of the datasets at the level of the individual datasets or combinations thereof. However, the statistical properties of data may vary depending on the objective to be achieved. For example, while dynamic movement data are needed to identify humans moving in a field, the same data may not have the appropriate properties to identify a specific human among others or provide the detailed information required for a surgical operation.

The legislator requires that the purpose of the datasets be defined, focusing on characteristics or elements that are relevant to the specific geographical, contextual, behavioural, or functional environment in which the high-risk AI system is intended to be deployed. For instance, in the context of developing an AI system for autonomous agricultural machinery, the legislator requires the definition of the purpose of the datasets, focusing on characteristics or elements relevant to the specific geographical, contextual, behavioural, or functional environment in which the high-risk AI agricultural system is intended to be deployed.

The AI developer is responsible for creating and curating datasets that accurately capture the unique environmental and operational characteristics of the agricultural settings in which the autonomous machinery will be deployed. This includes factors such as the geographical and contextual environment. The datasets should capture geographical features such as terrain types, soil compositions, elevation variations, and weather patterns specific to the intended deployment

locations. This information is crucial for enabling the AI system to adapt to diverse agricultural landscapes and environmental conditions.

Furthermore, the datasets are required to encompass contextual factors relevant to agricultural operations, including field layouts, crop types, planting densities, and irrigation infrastructure. Understanding these contextual elements is essential for the AI system to make informed decisions based on the specific agricultural context.

Additionally, the datasets need to reflect the functional aspects of agricultural machinery operations, including equipment specifications, operational constraints, and safety protocols. These functional elements are critical for training the AI system to perform its tasks efficiently and safely in the agricultural environment.

By defining the purpose of the datasets in this manner, the AI developer ensures that the training data accurately represents the unique geographical, contextual, behavioural, and functional environment in which the high-risk AI system will operate. This tailored approach to dataset definition aligns with the legislator's requirement to capture the specific characteristics and elements relevant to the intended deployment environment.

6.9.1 Recordkeeping

Developers must maintain detailed documentation of the AI system's development process, including data sources, methodologies, and decision-making processes. This ensures transparency and facilitates oversight.

The recordkeeping requirements for high-risk AI systems include the automatic recording of events (logs) throughout the system's lifetime. These logs should enable the identification of situations that may pose a risk, facilitate post-market monitoring, and monitor the operation of high-risk AI systems. For high-risk AI systems, the logging capabilities should include recording the period of each use of the system, the reference database used for input data verification, the input data that led to a match, and the identification of the natural persons involved in result verification.

6.9.2 Transparency

High-risk AI systems must be transparent in their operations. Users should be informed about the system's capabilities, limitations, and the manner in which it processes data. Article 13 of the AI Act emphasizes the transparency and provision of information to deployers of high-risk AI systems.

High-risk AI systems must be designed and developed to ensure transparent operation, enabling deployers to interpret the system's output and use it appropriately. High-risk AI systems must be accompanied by clear and concise instructions for use, provided in an appropriate digital format or otherwise, to ensure that the information is relevant, accessible, and comprehensible to users.

Overall, the transparency requirements in the AI Act aim to ensure that deployers have access to comprehensive and understandable information about high-risk AI systems, enabling them to interpret and use the systems effectively while complying with relevant obligations.

6.9.3 Human oversight

There must be appropriate levels of human oversight to monitor the AI system's operation and to intervene when necessary. This ensures that decisions made by AI systems can be reviewed and corrected by humans. Human oversight is also a requirement provided by the AI Act³⁰ for high-risk AI systems.

High-risk AI systems must be designed to allow effective oversight by natural persons during the system's use, including the use of appropriate human-machine interface tools.

³⁰ Art. 14 of the AI Act.

Human oversight aims to prevent or minimize risks to health, safety, or fundamental rights that may arise from the use of the AI system, especially when such risks persist despite other regulatory requirements. It aims at controlling aberrant results by empowering humans to interrupt AI activities. Oversight measures, commensurate with the risks and level of autonomy of the AI system, can be implemented through measures built into the system. An example thereof is the so-called "kill switch" aimed at interrupting AI operation in case of errors or unclear situations (e.g., possible failure to identify a human).

The AI system must be provided to the user in a way that enables natural persons assigned to human oversight to understand the system's capacities and limitations, monitor its operation, interpret its outputs, intervene in its operation, and override its outputs if necessary.

For high-risk AI systems, additional measures are required to ensure that no action or decision is taken by the deployer based on system identification alone, unless separately verified and confirmed by at least two natural persons with the necessary competence, training, and authority. Overall, the requirement outlines the effective human oversight of high-risk AI systems to mitigate risks and ensure responsible use.

6.9.4 Accuracy, robustness, and cybersecurity

AI systems must be robust, secure, and perform accurately under all conditions for which they are designed. This includes ensuring resilience to attacks and attempts at manipulation.

High-risk AI systems must be designed to achieve an appropriate level of accuracy, robustness, and cybersecurity throughout their lifecycle. This provision is closely related to the new Cyber Resilience Regulation³¹ that aims to establish a high cyber resilience of systems deployed in the EU. The levels of accuracy and relevant accuracy metrics of high-risk AI systems must be declared in the accompanying instructions for use.

High-risk AI systems must be resilient against attempts by unauthorized third parties to alter their use, outputs, or performance, by exploiting system vulnerabilities. Technical solutions for ensuring cybersecurity should be appropriate to the relevant circumstances and risks. These solutions should address AI-specific vulnerabilities, including measures to prevent, detect, respond to, resolve, and control attacks aimed at manipulating training datasets, pre-trained components, inputs, confidentiality, and model flaws.

6.9.5 Technical documentation

Providers of high-risk AI are required to provide technical documentation as a prerequisite to placing the AI products on the market.

The technical documentation will be required to be provided to the competent authorities and notified bodies for assessing the compliance level of the AI systems with the requirements of the AI Act. Annex IV³² provides basically a checklist of the documents to be provided, with some exceptions and simplifications to be applied to small- and medium-sized enterprises.

6.10 Sandboxes

The legislator has introduced the provisions for testing high-risk AI systems in real-world conditions, providing mandatory sandboxes for high-risk AI.

The AI regulatory sandbox allows for the processing of personal data collected for other lawful purposes for the development, training, and testing of certain AI systems under specific conditions. AI systems developed in the sandbox are intended to safeguard substantial public interests in areas such as public safety, public health, environmental protection, energy sustainability, transport

³¹ <https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=CELEX:52022PC0454>

³² A copy of the checklist of Annex IV is provided at the end of the document.

systems, public administration, and public services. Effective monitoring mechanisms and response mechanisms must be in place to identify and mitigate any high risks to the rights and freedoms of data subjects during sandbox experimentation.

6.11 Requirements for general purpose AI (e.g., transformer models)

" 'General purpose AI model' means an AI model, including when trained with a large amount of data using self-supervision at scale, that displays significant generality and is capable to competently perform a wide range of distinct tasks regardless of the way the model is placed on the market and that can be integrated into a variety of downstream systems or applications. This does not cover AI models that are used before release on the market for research, development and prototyping activities."³³

The AI Act outlines transparency obligations for providers and users of certain AI systems and general purpose AI (GPAI) models³⁴, as well as the classification of GPAI models with systemic risk. However, new developments may lead to also applying pre-trained models to different applications, such as autonomous machine operations in the field.

GPAI models are classified as having systemic risk if they have high-impact capabilities or are embedded in a product that may lead to a high-risk classification. The Commission may classify a GPAI system as high risk if it is determined to present high-impact or equivalent capabilities, based on a qualified alert issued by the scientific panel. A GPAI model is presumed to have high-impact capabilities when the cumulative amount of compute used for its training, measured in floating point operations (FLOPs), is greater than 10^{25} .

The Commission is empowered to adopt delegated acts to amend the thresholds and supplement benchmarks and indicators in light of evolving technological developments. These provisions aim to ensure transparency in the use of AI systems and models and to classify GPAI models with systemic risk based on their impact capabilities.

6.12 Certification

The AI Act provides that high-risk AI systems need to submit to a fundamental rights impact assessment, in the form of third-party (notified body) certification.

Prior to deploying a high-risk AI system, deployers must perform an assessment of the impact on fundamental rights that the use of the system may produce. This assessment includes a description of the deployer's processes, the intended use period and frequency, the categories of affected individuals, specific risks of harm, implementation of human oversight measures, and measures to be taken in case of materialized risks. If any factors change during the use of the system, the deployer must update the information.

Once the impact assessment is performed, the deployer must notify the market surveillance authority of the results, submitting a filled template as part of the notification. There are exemptions in certain cases.

Overall, the fundamental rights impact assessment is a crucial step in ensuring that the deployment of high-risk AI systems considers the potential impact on fundamental rights and provides appropriate measures to address any identified risks.

Furthermore, the AI Act provides the requirement of the registration of high-risk AI systems at the local or European AI Authority. Once classified as high-risk, AI systems need to be registered in an EU database, making information about these systems publicly available and ensuring accountability.

³³ Art. 2 44(b) AI Act.

³⁴ Art. 52 AI Act.

The registration may contain all technical information as collected under the technical information annex.

6.13 Penalties for non-compliance

Noncompliance with the requirements of the AI Act is subject to penalties, independently of whether the AI has created damage or not.

"Noncompliance with the prohibition of the artificial intelligence practices is subject to administrative fines of up to 35 000 000 EUR or, if the offender is a company, up to 7% of its total worldwide annual turnover for the preceding financial year"³⁵ and non-compliance of an AI system for other infringements shall be subject to administrative fines of up to 15 000 000 EUR or, if the offender is a company, up to 3% of the worldwide annual turnover. Providing incorrect, incomplete or misleading information is punishable with a fine up to 7 500 000 EUR or 1% of the worldwide turnover.

However, for the actual fine, each case is evaluated on case-by-case basis. Other factors, like the nature, duration and gravity of the infringement, or cooperation with the authorities, may play a role in the calculation of the fine.

6.14 Code of conduct

The AI Act outlines the encouragement and facilitation of codes of conduct for the voluntary application of specific requirements to AI systems. These codes should take into account available technical solutions and industry best practices. The codes of conduct concerning the voluntary application of specific requirements to all AI systems are ideally based on clear objectives and key performance indicators. This includes, for example but not only, promoting the ethical development and use of AI and minimizing the risks of AI systems, facilitating an inclusive and diverse design of AI systems, and assessing and preventing negative impacts.

Codes of conduct may be developed by individual providers or deployers of AI systems, the organizations representing them, or a combination of both, with the involvement of interested stakeholders and their representative organizations, including civil society organizations and academia.

7 UN Resolution on AI

In March 2024, the United Nations General Assembly adopted a groundbreaking Resolution aimed at directing the use of AI for global benefit. The Resolution, titled "Seizing the opportunities of safe, secure, and trustworthy artificial intelligence systems for sustainable development", was adopted without a vote. It emphasizes the need to address racial discrimination worldwide, including through reparations. The Assembly resolved to bridge the AI and other digital divides between and within countries and promote safe, secure, and trustworthy AI systems, to accelerate progress towards the full realization of the 2030 Agenda for Sustainable Development [b-UN].

The Resolution focuses on the potential of safe, secure, and trustworthy AI systems for sustainable development. It begins by reaffirming international law and recalling various Resolutions related to sustainable development, human rights, and the impact of technological advancements on achieving the Sustainable Development Goals. The document emphasizes the importance of AI systems being human-centric, reliable, ethical, and promoting human rights, while also recognizing the potential risks associated with their improper or malicious use.

The Resolution calls for bridging digital divides, promoting inclusive access to the benefits of AI systems, and fostering an enabling environment for innovation and entrepreneurship. It emphasizes the need for effective partnerships, regulatory and governance approaches, and capacity building,

³⁵ Art 71 of the AI Act.

particularly for developing countries. The document outlines a series of actions for Member States and stakeholders, including enhancing digital infrastructure connectivity, promoting human rights and fundamental freedoms throughout the AI life cycle, and safeguarding privacy and personal data.

Furthermore, it highlights the significance of data governance, international cooperation, and collaboration among public and private sectors, academia, and research institutions in promoting safe and trustworthy AI systems. The document also encourages the development of effective safeguards, risk management mechanisms, and impact assessments throughout the AI life cycle, while promoting the transparency, predictability, and understandability of AI systems' decisions. Additionally, it stresses the importance of addressing gender and digital divides, supporting digital training, and enhancing access to AI benefits, especially in developing countries.

The Resolution calls upon specialized UN agencies and related organizations to assess and enhance their response to AI opportunities and challenges, emphasizing the need to close digital divides and promoting inclusive international cooperation. The document acknowledges the unique role of the UN system in reaching global consensus on safe and trustworthy AI systems, consistent with international law, human rights, and sustainable development goals, while promoting inclusive international cooperation and the representation of developing countries.

In the context of this Resolution, safe, secure, and trustworthy AI can significantly benefit digital agriculture by enhancing efficiency, productivity, and sustainability. AI technologies can analyse vast amounts of agricultural data, such as weather patterns, soil conditions, and crop health, to provide valuable insights for farmers. This can lead to optimized resource allocation, improved crop management, and better decision-making processes. Additionally, AI-powered predictive analytics can help farmers anticipate and mitigate potential risks, such as pests, diseases, and adverse weather conditions, thereby improving crop yields and reducing losses.

Furthermore, AI can enable precision agriculture by facilitating the use of autonomous vehicles, drones, and robotic systems for tasks such as planting, irrigation, and harvesting. These technologies can operate with high precision, reducing resource wastage and environmental impact while increasing overall efficiency. Additionally, AI can support the development of smart farming systems that integrate data from various sources to automate and optimize agricultural processes.

8 AI standards

High-risk AI systems must comply with existing regulatory standards, including those related to privacy, non-discrimination, and consumer rights, ensuring the protection of fundamental rights.

Several institutions are working to assess standards requirements for AI, on a horizontal level targeting AI development criteria and, on a vertical level, targeting specific sectors.

While horizontal standards are conceived to be applied universally across diverse AI applications, offering a consistent framework for addressing common aspects such as data privacy, security, and ethical considerations, simplifying compliance and ensuring compatibility of AI systems across sectors and industries, vertical standards target the unique requirements and challenges of sector-specific AI applications, ensuring relevance and effectiveness in addressing domain-specific needs.

Horizontal and vertical standards provide their respective advantages and challenges. The key lies in finding the right balance and harmonization between horizontal and vertical standards, to ensure a comprehensive and effective approach to AI standardization.

Furthermore, standards may be applicable to the design and development phase of AI or may target the final results. While standards for the design and development phase focus on building quality, ethicality, and risk mitigation into AI systems from the outset, results-focussed standards are aimed at ongoing validation and improvement.

High-risk AI systems that are in conformity with harmonized standards or parts thereof, the references for which have been published in the Official Journal of the European Union, in accordance with Regulation (EU) 1025/2012, shall be presumed to be in conformity with the requirements.

Organizations that develop international standards, like the International Telecommunication Union (ITU), play a significant role in shaping global standards and governance frameworks for emerging technologies, including AI. ITU develops technical standards and guidelines that facilitate the interoperability, compatibility, and reliability of AI systems. These standards cover various aspects of AI, including data formats, algorithms, performance metrics, and ethical considerations. By establishing common standards, ITU promotes the responsible development and deployment of AI technologies on a global scale. ITU standards often incorporate ethical principles to guide the development and use of AI systems. This includes considerations such as fairness, transparency, accountability, and privacy protection. By embedding ethical norms into technical standards, ITU aims to foster trust and confidence in AI technologies among governments, businesses, and society at large. Various technical groups within ITU-T, including [ITU-T Study Group 13](#), [ITU-T Study Group 16](#) and [ITU-T Study Group 20](#), have developed several standards revolving around the implementation of AI across different verticals. In the context of agriculture, ITU-T has developed various standards such as ITU-T Y.4482 – *Requirements and framework for smart livestock farming based on the Internet of things*. The ITU/FAO Focus Group on "Artificial Intelligence (AI) and Internet of Things (IoT) for Digital Agriculture" (FG-AI4A) also develops best practices and underscores obstacles associated with the utilization of AI and IoT-based technologies in the agricultural domain³⁶.

9 Ethical aspects of AI in agriculture

AI can analyse statistical patterns in large datasets to provide insights that would not be gained through classical or traditional means, due to the sheer volume of data it is capable of analysing in such an efficient manner. However, with this immense statistical power, the use of AI raises questions around how to utilize its vast array of tools in an ethical manner.

Independently from the actual narrow or general scope of the AI application, the combination of factors such as data, models, parameters, and activations, as well as training, represent the factors that may lead to the need for ethical assessment.

If the AI inputs are biased, incomplete or qualitatively insufficient, the resulting output may provide a statistical bias or lead to an aberrant result. Furthermore, where the amount of data enters the billions or trillions of parameters, it becomes proportionally more difficult to identify biased data or the process that led to a biased result.

Ethical and social questions may also arise from the logic of the chosen models and parameters³⁷. A model that prioritizes the safety and protection of users and bystanders may have a chance of success if applied by all manufacturers; however, ethical challenges could arise if such models are only applied by a part of the market, while other manufacturers choose a more economic approach that does not prioritize safety and protection^{38, 39}.

³⁶ For more information on AI-based standards related to agriculture, see the FG-AI4A Report: Standardization gaps and roadmap for AI and IoT in digital agriculture.

³⁷ F. Poszler/M. Geißlinger, AI and Autonomous Driving: key ethical considerations (2021) [ResearchBrief_February2021_AutonomousVehicles_FINAL.pdf \(tum.de\)](#).

³⁸ M. Geißlinger/F. Poszler/C. Lütge, M. Lienkamp, Autonomous Driving Ethics: from Trolley Problem to Ethics of Risk (2021).

³⁹ N.Godall Ethical Decision Making During Automated Vehicle Crashes (2014).

The AI Act differentiates risk levels between less risky and high-risk applications, requiring the latter to meet higher legal standards⁴⁰ before being introduced in products for the market. Full ethical consideration of AI is difficult given the fast advancement and versatility of the technology and the different fields of application. Ethical analysis may vary based on the technical level of the specific use case of the AI. For example, a self-driving AI may require less consideration of fundamental rights, health or safety when operating on a remote field or in a secured enclosed space where interaction with people is not possible, compared to a machine operating in a crowded or unsecured environment.

AI may also require a balance of interests between the pros and cons of the use of AI. Agriculture may provide several examples of the innovation and enhancement potential in numerous fields of application.

Industrialization, including the movement of vast populations from rural to urban areas, has impacted agriculture, requiring the sector to modernize and increase production. This poses considerable challenges for sustaining population growth and satisfying the increased demand in food quality and supply. AI may enhance precision agriculture by optimizing farming operations through data-driven insights. AI can optimize machinery settings, and help farmers make better decisions, by providing advanced crop management capabilities through analysing historical data on crop yields, weather patterns, soil quality and geographic features. This analysis, coupled with farmers' experience, can drive improved overall results.

AI may also be used to analyse data on pest and crop diseases to predict where they may occur and how to prevent them, thereby minimizing crop damage and loss. AI may support efforts to achieve sustainability objectives by fostering rational resource usage, such as the balanced use of agricultural land and water resources and/or limiting environmental impacting factors by optimizing the use of nutrients and crop protection. Finally, AI can assist efforts to optimize economic resources, reducing costs, meeting demand, and preventing constraints.

While the use of AI has the potential to revolutionize the agriculture industry, its use may raise ethical questions regarding inequalities between producers and whether it is ethical to utilize or, rather, if it is still ethical not to use AI in agriculture to meet the demands of a growing global population.

Additionally, AI-based standards for agriculture also play a critical role in fostering trust, promoting responsible innovation, and addressing societal concerns associated with the adoption of AI technologies in farming. By adhering to these standards, stakeholders can harness the potential of AI to enhance productivity, sustainability, and resilience in agricultural systems, while upholding ethical principles and values. These standards present frameworks for clarifying responsibilities in case of AI system failures or unintended consequences. This includes establishing guidelines for liability attribution and ensuring that appropriate mechanisms are in place to address harm caused by AI technologies in agriculture. Furthermore, these standards will also promote transparency in the AI algorithms used in agriculture, ensuring that farmers understand how AI-based decisions are made. This includes providing explanations for recommendations or actions taken by AI systems to build trust and enable informed decision-making. Many international standards leveraging AI in agriculture will also aim to promote sustainable practices, by encouraging the development of AI technologies that minimize environmental impact, optimize resource usage, and support eco-friendly farming practices.

10 Conclusion

The AI Act introduced by the EU has significant implications for the agricultural sector, which is increasingly relying on AI technologies for efficiency and sustainable improvements. The Regulation

⁴⁰ Arts. 10 and 11 of the Draft AI Act COM/2021/206 final/UNECE Framework document on automated/autonomous vehicles (2019) WP.29-177-19.

categorizes AI systems into different risk levels, with specific requirements for those deemed high risk. In agriculture, AI applications such as precision farming, crop monitoring, and automated machinery could fall under these Regulations, depending on their potential impact on safety and fundamental rights.

For the agricultural sector, the AI Act necessitates a careful evaluation of AI technologies, to ensure they comply with the established standards for safety, transparency, and ethical use. This could mean additional regulatory compliance efforts for developers and users of AI in agriculture, particularly for systems classified as high risk. These systems may require thorough risk assessments, adherence to strict data governance practices, and the implementation of robust human oversight mechanisms.

The AI Act aims to promote the development of secure and trustworthy AI, which could enhance innovation in agriculture by providing a clear regulatory framework. This could encourage investment in AI technologies that offer solutions for sustainable farming practices, increased productivity, and environmental protection. However, the need for compliance could also pose challenges, especially for smaller agricultural businesses and startups, which may face resource constraints in meeting the regulatory requirements.

AI-based standards for agriculture are instrumental in building trust, fostering responsible innovation, and addressing societal concerns related to the adoption of AI in farming. Adhering to these standards allows stakeholders to harness AI's potential to improve productivity, sustainability, and resilience in agricultural systems, while upholding ethical principles. These standards establish frameworks for clarifying responsibilities in cases of AI system failures or unintended consequences, including guidelines for liability attribution and mechanisms to address harm caused by AI technologies in agriculture. Furthermore, these standards promote transparency in the AI algorithms used in agriculture, by ensuring that farmers understand the basis of AI-driven decisions. This involves providing explanations for AI recommendations or actions to build trust and enable informed decision-making.

Bibliography

- [b-EU AI] Artificial Intelligence Act.
- [b-UN] Seizing the opportunities of safe, secure and trustworthy artificial intelligence systems for sustainable development.
-